

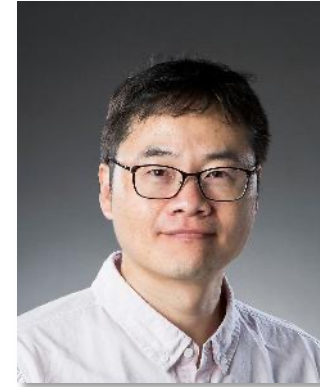
HInT: Hypergraph Infusion at the Structural Layers Improves Table Understanding



Wonjin Lee



Soomi Jeong



Kwang In Kim

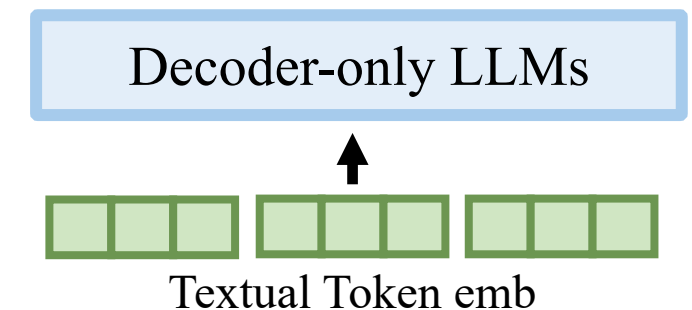
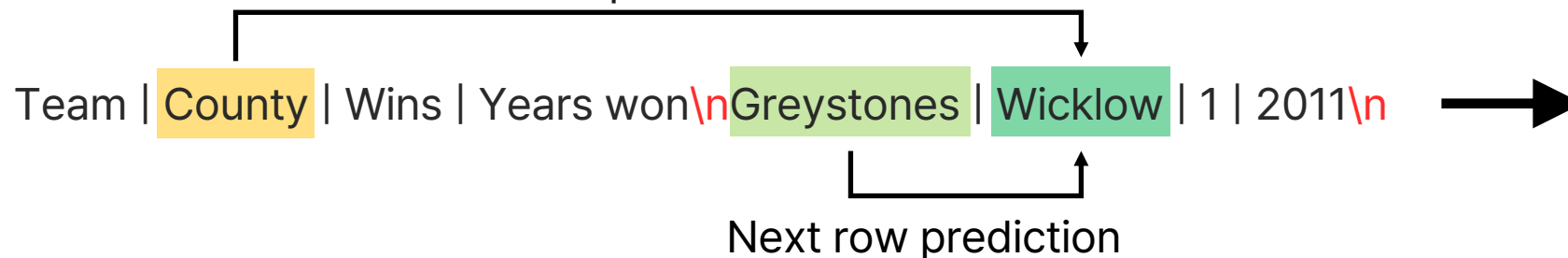
Key challenges: Loss of structural information

Linearizing tables obscures row/column relations and higher-order structure

Team	County	Wins	Years won
Greystones	Wicklow	1	2011
Ballymore Eustace	Kildare	1	2010
Maynooth	Kildare	1	2009

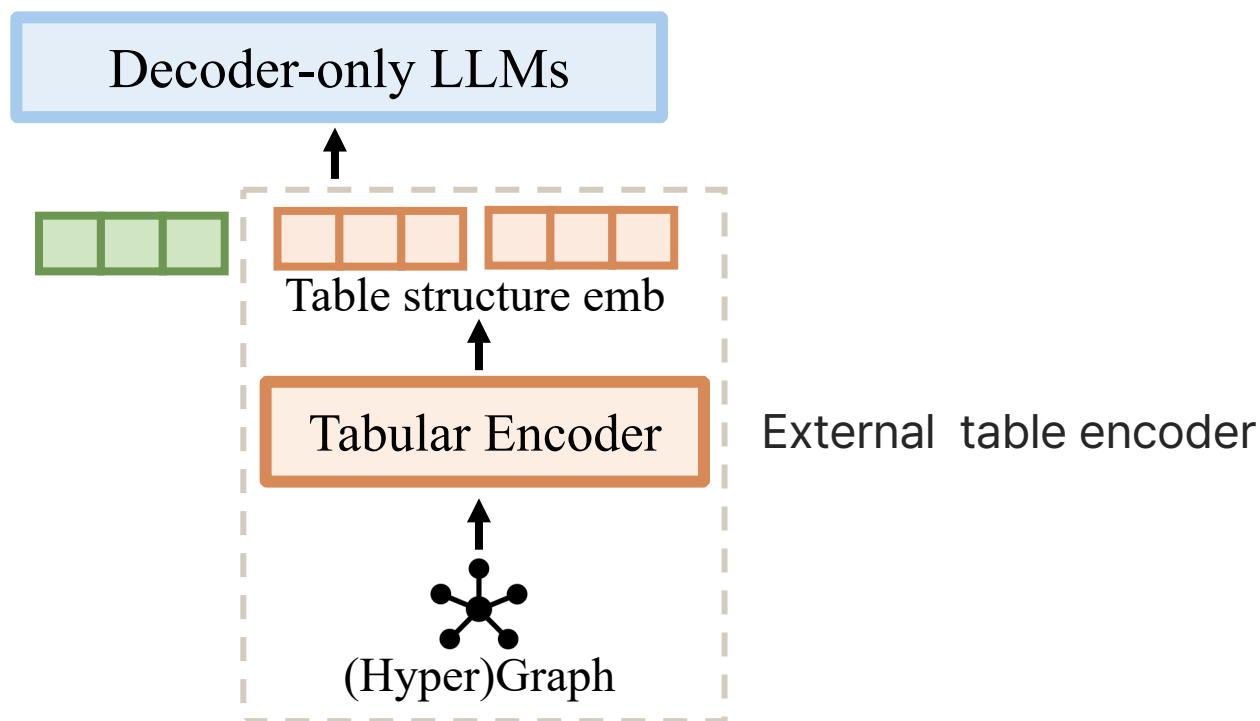
↓ Linearized table

Next column prediction



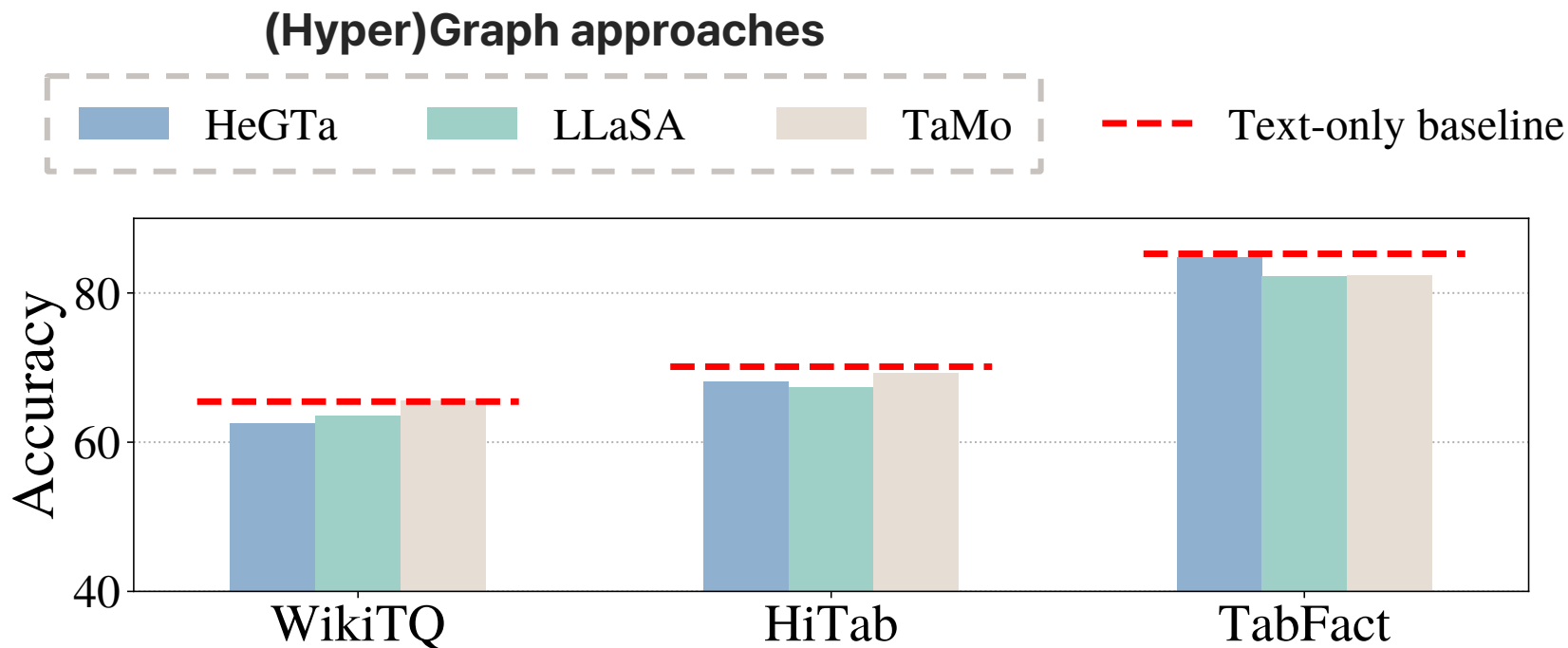
Prior works: Table encoder for structure modeling

Table structure is encoded separately and injected as structure-aware features



Prior works: Limitations of external table encoders

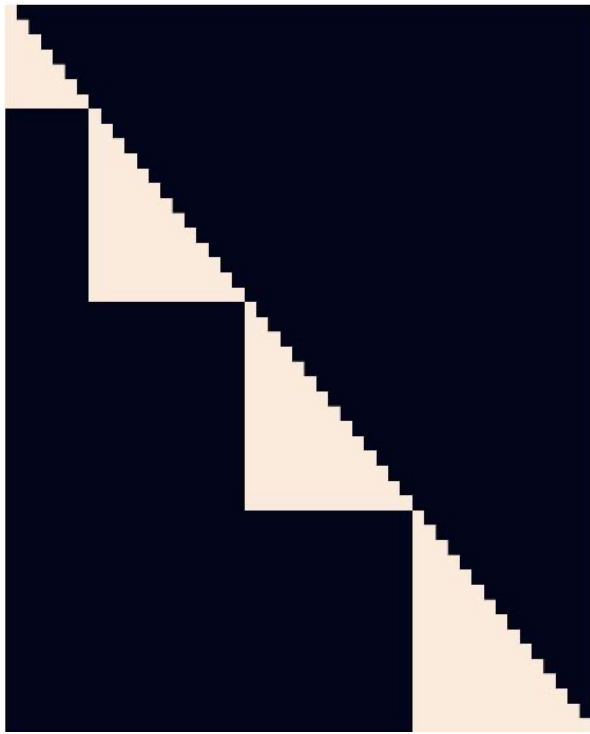
(Hyper)Graph approaches yields comparable, and in some cases lower performance than text-only baselines under multi-step reasoning setting



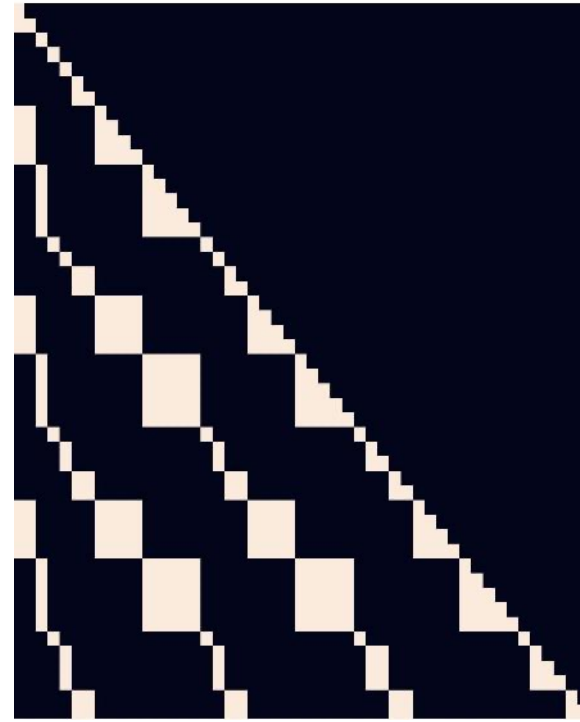
Key message: The bottleneck is not just extracting structure, but integrating it where reasoning happens inside the LLM

Our approach: Identifying structurally significant attention heads

Table structure is encoded as row-wise and column-wise attention masks



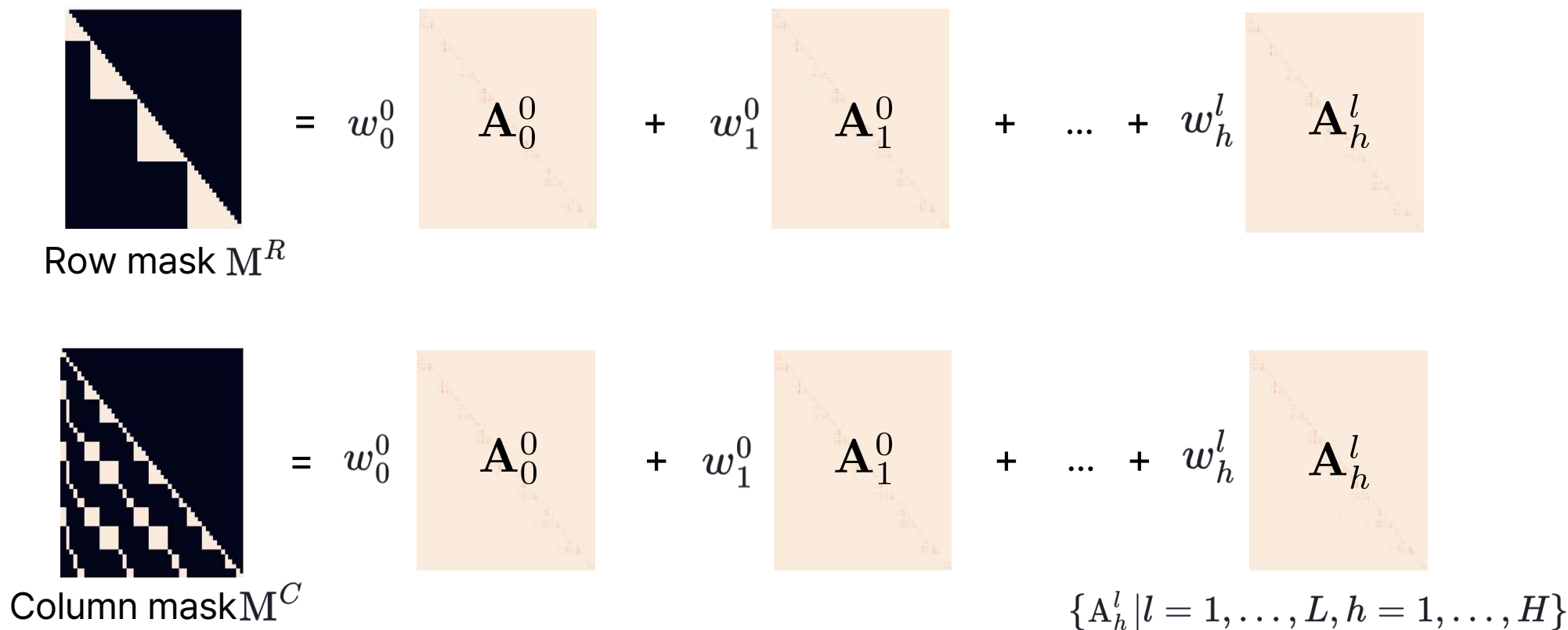
Row mask M^R



Column mask M^C

Our approach: Identifying structurally significant attention heads

We measure the alignment between
attention patterns and the structural masks



The diagram illustrates the decomposition of structural masks into weighted attention patterns. It consists of two rows of equations. The top row shows a 'Row mask M^R ' as a square matrix with a dark blue upper triangle and a light beige lower triangle. This is equated to a sum of weighted attention patterns: $w_0^0 \mathbf{A}_0^0 + w_1^0 \mathbf{A}_1^0 + \dots + w_h^l \mathbf{A}_h^l$. Each $\mathbf{A}_h^l$ is shown as a square matrix with a light beige background and a diagonal of small red arrows pointing from top-left to bottom-right. The bottom row shows a 'Column mask M^C ' as a square matrix with a dark blue lower triangle and a light beige upper triangle. This is equated to a similar sum: $w_0^0 \mathbf{A}_0^0 + w_1^0 \mathbf{A}_1^0 + \dots + w_h^l \mathbf{A}_h^l$. Below the second equation, the set notation $\{\mathbf{A}_h^l | l = 1, \dots, L, h = 1, \dots, H\}$ is provided.

$$\text{Row mask } M^R = w_0^0 \mathbf{A}_0^0 + w_1^0 \mathbf{A}_1^0 + \dots + w_h^l \mathbf{A}_h^l$$
$$\text{Column mask } M^C = w_0^0 \mathbf{A}_0^0 + w_1^0 \mathbf{A}_1^0 + \dots + w_h^l \mathbf{A}_h^l$$

$\{\mathbf{A}_h^l | l = 1, \dots, L, h = 1, \dots, H\}$

Our approach: Identifying structurally significant attention heads

We measure the alignment between attention patterns and the structural masks

$$\mathcal{L}(\mathbf{w}) = \left\| \mathbf{m}^{\text{R}} - \sum_{l,h} w_h^l \mathbf{a}_h^l \right\|_2^2 + \lambda \|\mathbf{w}\|_1,$$

$$\mathbf{a}_h^l = \text{vec}(A_h^l), \quad \mathbf{m}^{\text{R}} = \text{vec}(M^{\text{R}}), \quad \mathbf{w} = \{w_h^l\}_{l,h}.$$

$$\{A_h^l | l = 1, \dots, L, h = 1, \dots, H\}$$

We ranked the attention heads by the learned LASSO coefficients

Our approach: Identifying structurally significant attention heads

Removing the top ranked heads results in a larger performance drop than removing randomly selected heads

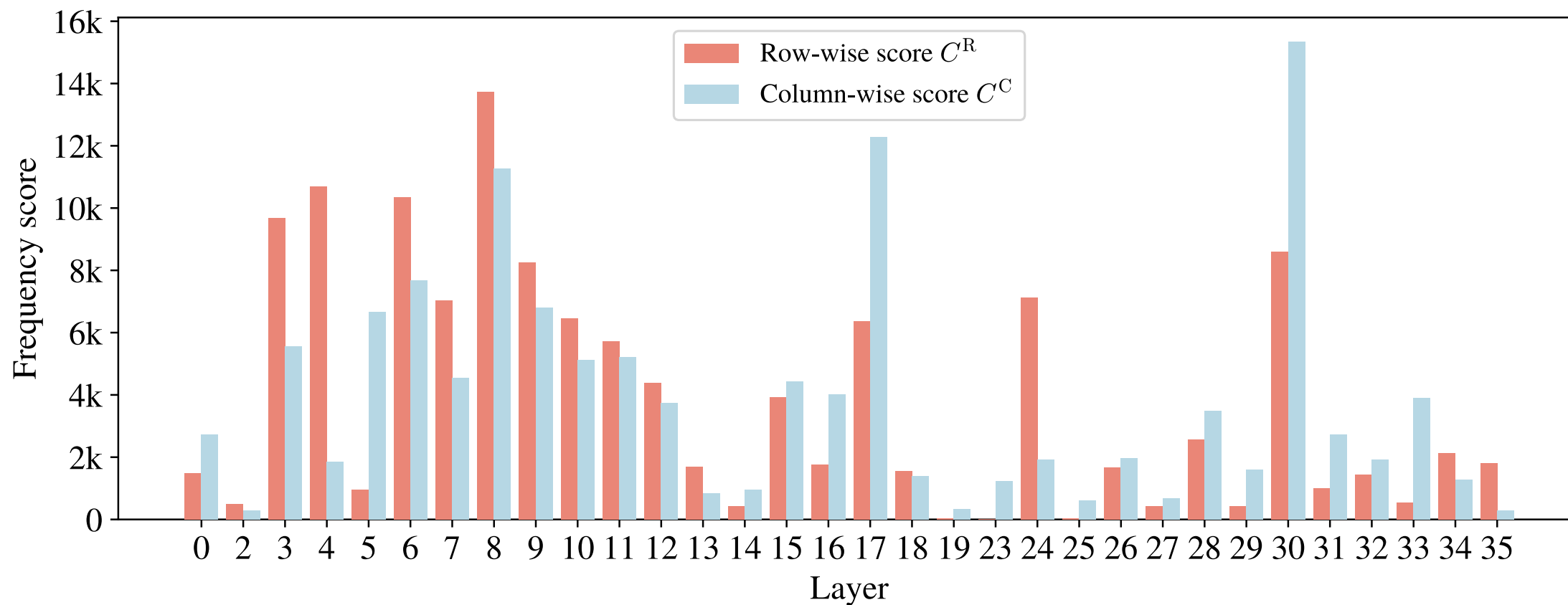
	Original	Ablate structural heads	Ablate random heads
Qwen 2.5-3B	47.51	12.73	32.29

(in % exact match; EM)

Our approach: Identifying structural layers

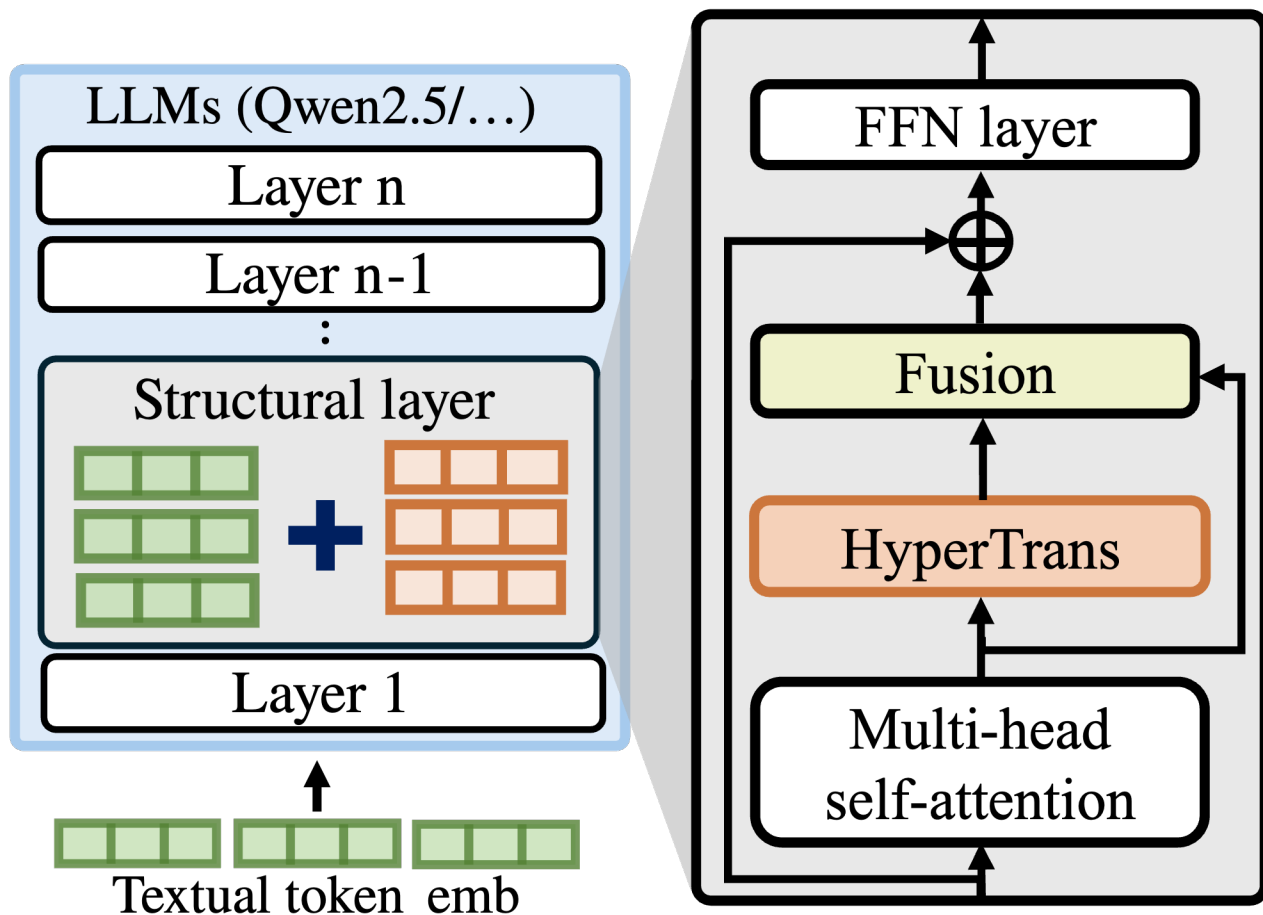
Structural layers (top-12)

$$\mathcal{L}^* = \{3, 4, 5, 6, 8, 9, 10, 12, 15, 17, 24, 30\}$$



Our approach: Injecting structural information at structural layers

Converting token representations into node and hyperedge features



Our approach: Injecting structural information at structural layers

Given a table, we construct a hypergraph: $\mathcal{G} = (\mathcal{V}, \mathcal{E})$

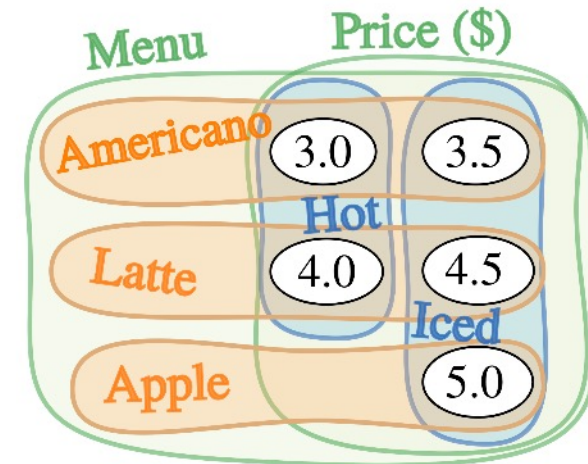
$\mathcal{V} = \{1, \dots, N_c\}$: table cells

$\mathcal{E} = \{e_1, \dots, e_{N_h}\}$: hyperedges

$\mathcal{E} = \{\text{Menu}, \text{Price } (\$), \text{Americano}, \text{Latte}, \text{Apple}, \text{Hot}, \text{Iced}\}$

Menu	Price (\$)	
	Hot	Iced
Americano	3.0	3.5
Latte	4.0	4.5
Apple		5.0

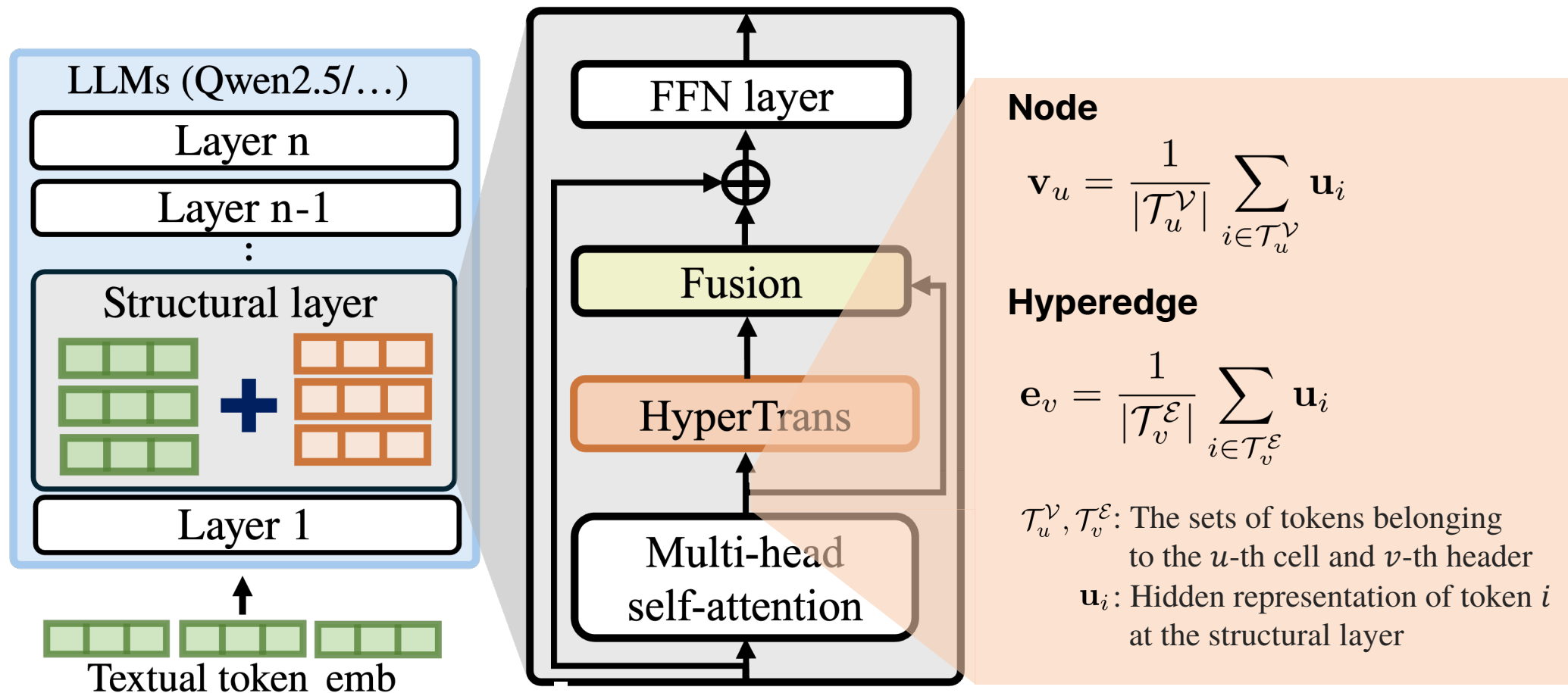
Table



Hypergraph

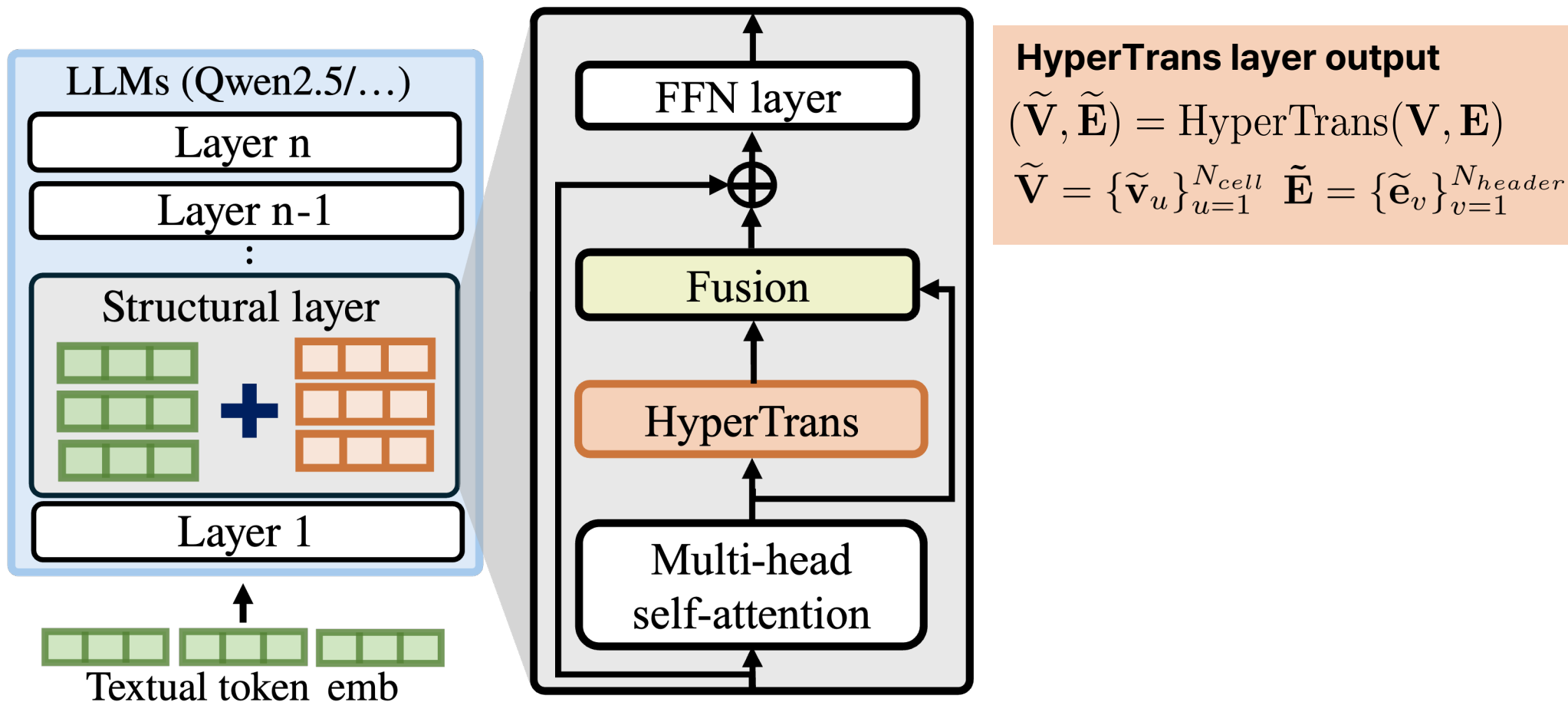
Our approach: Injecting structural information at structural layers

Converting token representations into node and hyperedge features

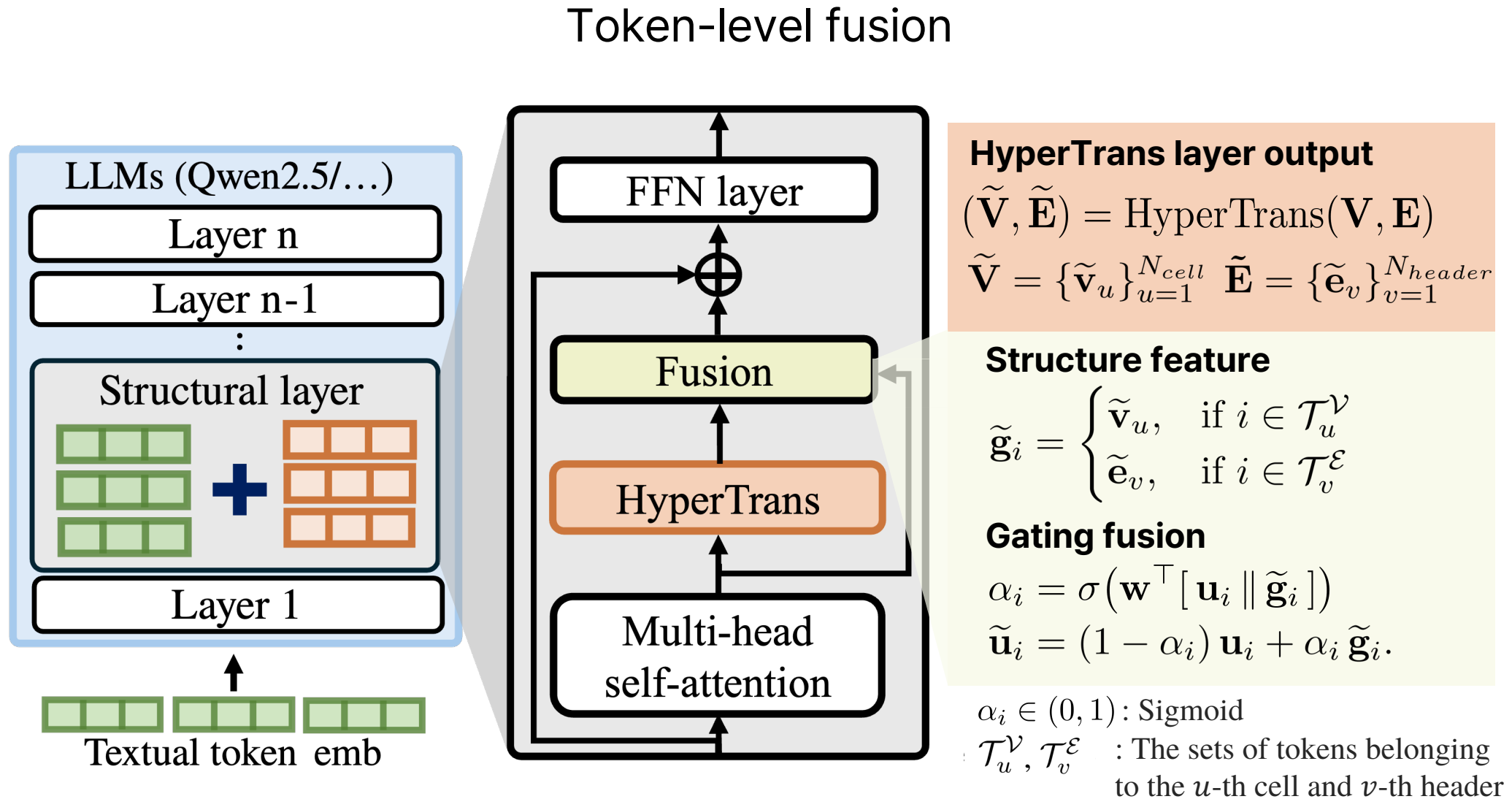


Our approach: Injecting structural information at structural layers

Hypergraph message passing



Our approach: Injecting structural information at structural layers



Results: Structural layers are the most effective injection points

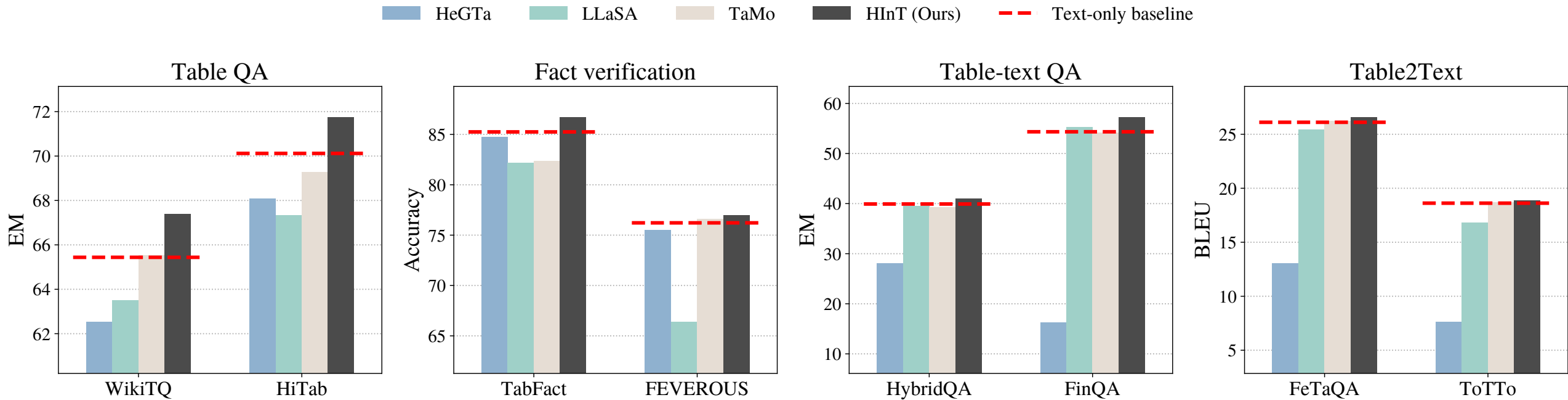
Consistently outperforms random or full-layer injection

Model: Qwen2.5-3B-Instruct

	Random	Full	Structural Layers
WikiTQ	65.71	66.37	67.22

(in % exact match; EM)

Results: HInT achieves the best overall performance



Thank For Listening

